



Operationalising Leadership Architecture into a Talent- Management Computation Instrument for Indonesia's Immigration and Corrections Ministry

Mochammad Sofyan Arief^{1*}

m.sofyan.arief@gmail.com

Head of Data and Information Systems, Human Resources Development
Agency for Immigration and Corrections, Ministry of Immigration and
Corrections, Jakarta, Indonesia

*Corresponding Author: Mochammad Sofyan Arief

E-mail: m.sofyan.arief@gmail.com

ABSTRACT

Leadership-architecture frameworks in public administration typically specify what officials ought to do at each level, without specifying how to compute it, so such frameworks work as development maps but not as ranking instruments for succession. This article addresses that gap for the Indonesian Ministry of Immigration and Corrections. It converts the three-component leadership architecture of Minister of Finance Regulation 191/PMK.01/2018, namely leadership capability, work values and time management, into a computation instrument anchored in the national talent-management regime of National Civil Service Agency Decree 411/2025. The instrument comprises five sequential layers, twelve equations, complete conversion scales, and weight matrices differentiated by target-position level and type, producing a Final Talent Score for a specified officer-position pair. A tiered institutional parameter is proposed to correct a defect in the baseline design whereby leadership time management loses influence as the hierarchy is ascended. The instrument is audited against Star ASN, the Ministry's personnel-record application covering 65,412 employees and 8,353 structural positions. Existing records support between 46.25 and 66.00 per cent of the Final Talent Score depending on target level, with the shortfall concentrated in four indicators requiring new assessment modules. A worked simulation returns Final Talent Scores of 88.34, 87.54 and 86.33 for one officer against three target-position types, and sensitivity analysis shows differentiated weighting discriminates only among officers with lopsided certification profiles. A web-based reference architecture and governance rules are proposed.

Keywords: Civil Service Reform, Competency Measurement, Leadership Architecture, Succession Planning, Talent Management

INTRODUCTION

Frameworks that describe leadership in tiers are among the most widely adopted instruments in public-sector human resource management (Asif & Rathore, 2021; Nzimande et al., 2025). Bornay-Barrachina et al. (2025) believed that their appeal is intuitive: if the demands placed on a first-line supervisor differ in kind, and not merely in degree, from those placed on the head of an organisation, then development and appointment should be governed by explicit, level-specific statements of required capability. The leadership-pipeline tradition established this logic in the corporate sector, a foundational framework whose central premise the field continues to test rather than discard (Charan et al., 2011). A recent meta-analysis synthesising more than one hundred and fifty studies confirms that administrative leadership in the public sector produces measurable effects on organisational and employee outcomes, although the strength of those effects varies with context and method (Backhaus & Vogel, 2022).

A recurring weakness of such frameworks, however, is that they specify content without specifying computation. They state what capabilities matter and how they escalate, but not how a given officer is to be scored against them, how the resulting scores are to be weighted, or where the underlying data are to be obtained (Koedijk et al., 2021). This gap is not merely technical. Where a framework cannot generate a defensible number, it cannot be used to rank candidates, and appointment decisions revert to the discretionary judgement the framework was introduced to discipline. Talent-management scholarship continues to note that the field suffers from precisely this deficit of operationalisation: a review of one hundred and ninety-two studies finds that research on talent still adopts an unresolved, binary conceptualisation of the construct rather than a measurable one, and public-sector evidence from the United Arab Emirates shows that the link between talent-management practice and individual work performance depends heavily on mechanisms, such as line managerial support, that frameworks rarely specify (Semaihi et al., 2023; Vardi & Collings, 2023). The same operationalisation gap appears in Indonesian public-sector scholarship, where a study of organisational citizenship behaviour as a mediator between commitment, satisfaction and performance treats these constructs as psychometric variables to be measured rather than as normative aspirations to be asserted (Budiyanto, 2025).

The Indonesian Ministry of Immigration and Corrections (Kementerian Imigrasi dan Pemasyarakatan, hereafter the Ministry) presents an unusually demanding instance of this problem. Established as a separate ministry in 2024, it discharges two distinct state functions, immigration administration and the corrections system, through 65,412 employees occupying 8,353 structural positions distributed across roughly 900 work units nationwide. The distribution is steeply pyramidal: 4,738 positions (56.7 per cent) sit at Echelon V and 2,602 (31.2 per cent) at Echelon IV, so that approximately seven in eight structural leaders are first- or

second-line supervisors in direct contact with inmates, service applicants and foreign nationals. Provision for their development is inversely proportional to their number. Of the 8,353 incumbents, only 1,664 (20 per cent) have completed the tiered leadership training corresponding to their level; completion reaches 95 to 100 per cent at Echelons I and II but falls to 6 per cent at Echelon V.

Two regulatory developments frame any attempt to address this. First, the framework the Ministry had intended to adapt has been operationally superseded at its source. Minister of Finance Regulation 191/PMK.01/2018 introduced a three-component Leadership Architecture, comprising leadership capability, leadership work values and leadership time management, across five leadership stages. In June 2025, however, Minister of Finance Regulation 38 of 2025 reorganised the entire preparation and filling of positions around talent management. Regulation 191/2018 was not formally revoked and still appears in the preamble of the later regulation, but every operational function once attached to it has been assumed by the talent-management regime (Ererdi et al., 2022). Its residual value therefore lies in its conceptual framework, not in its mechanism. Second, Decree of the Head of the National Civil Service Agency Number 411 of 2025 obliges every Indonesian government institution to implement talent management using nationally uniform parameters, weights and assessment mechanisms, and to do so through a web-based application integrated with national personnel systems. This mandate follows a broader domestic pattern: Indonesian public-service agencies that have digitalised core administrative functions report measurable gains in service-delivery efficiency and accountability, alongside implementation gaps tied to institutional readiness (Sain et al., 2025).

This article treats that constellation as an opportunity rather than an obstacle. Its argument is that the conceptual value of a three-component leadership architecture can be preserved precisely by discarding its original measurement mechanism and re-anchoring each component in a measurement regime that is legally binding, nationally uniform and auditable. The contribution is threefold. It specifies a fully worked computation instrument, comprising twelve equations, complete conversion scales, and weight matrices differentiated by target-position level and type, that translates a normative leadership architecture into a single auditable score. It conducts a traceability audit establishing, quantitatively, what proportion of that score the institution's existing personnel data can actually support, and which indicators cannot yet be computed at all. And it proposes a reference component architecture for web-based implementation, together with the governance rules, namely weight versioning, audit trail, and separation of certification registers, without which a computed score cannot survive challenge.

LITERATURE REVIEW

The Three-Component Architecture and The Shift of Measurement Regime

The architecture introduced by Regulation 191/2018 rests on three components that answer three distinct questions about an official. Leadership capability asks what the officer is able to do, expressed as statements of technical, cross-functional, managerial and socio-cultural competency. Leadership work values asks what the officer has demonstrated, expressed as statements of character in relation to self, to others and to the work (Lušňáková et al., 2021). Leadership time management asks what the officer has invested, expressed as a shifting proportion of time between managing tasks and managing people as the hierarchy is ascended (Aeon et al., 2021).

The strength of the scheme is the clarity of its gradation; its weakness is the normative character of its mechanism. All three components are expressed as statements about what a leader ought to possess and do, unaccompanied by any weighted instrument capable of producing a comparable figure between individuals. The scheme functions well as a map of developmental direction and poorly as a basis for ranking in appointment. This is the same limitation that Boyatzis (1991) and Spencer and Spencer (2008) sought to overcome in competency modelling by insisting on behaviourally anchored, scoreable indicators, and the concern has not receded with time: research on leadership succession still identifies the absence of a measurable competency base, rather than the absence of a competency vocabulary, as the principal obstacle to building defensible succession pipelines (Jackson & Dunn-Jensen, 2021).

The third component deserves particular attention, because it is the one this article transforms most substantially. Time allocation between managing tasks and managing people is not independently measurable in the Ministry, for two reasons. Methodologically, the leadership content of time use is already absorbed into performance appraisal: under Government Regulation 30 of 2019, appraisal covers both work results and work behaviour, and the work-behaviour element of a structural post consists largely of coaching subordinates, building cooperation and mobilising teams. To measure time allocation separately is to count the same thing twice. Practically, no instrument records what proportion of a prison governor's working hours is spent coaching rather than executing, so any score would rest on self-report and could not be audited. This concern is consistent with a broader finding in the merit-system literature, that measurement designs which cannot be traced to verifiable evidence tend to erode rather than reinforce the government performance they are meant to support (Oliveira et al., 2024).

The National Measurement Regime as The Substituted Mechanism

Decree 411 of 2025 supplies what Regulation 191/2018 lacks. It establishes two orthogonal parameters. The performance parameter, forming the vertical axis of a nine-box talent map, comprises Core Performance weighted at 60 per cent and

Reinforcing Performance at 40 per cent. The potential parameter, forming the horizontal axis, comprises Competency at 40 per cent, Potential at 25 per cent, Qualification at 20 per cent, and Integrity and Morality at 15 per cent. Each indicator carries a defined conversion scale, and the two axis scores combine in equal proportion into a Talent Score. The Decree further permits institutions to conduct a technical competency assessment after nine-box mapping and to combine it with the Talent Score using weights that vary by target-position level, while leaving the internal composition of that technical assessment to institutional discretion.

Two features of this arrangement are decisive for the present design. The correspondence between the first two architecture components and the two axes is direct: capability is what the potential axis measures and demonstrated work value is what the performance axis measures (Le Fouest & Mulleners, 2024). And the discretion left over the composition of the technical competency assessment provides a legitimate aperture through which the third component may re-enter the computation in a measurable form. That this aperture must be exercised inside a nationally uniform, web-based application is itself consistent with recent scholarship on algorithmic governance in the public sector, which finds that digital government systems only deliver on their promise of consistency when the discretion left to individual institutions is bounded by rules that are transparent and auditable rather than implicit (Criado et al., 2025). The same logic underlies the requirement, developed later in this article, that a computed score be traceable to its evidentiary source: administrative decisions produced by a formula rather than a person are accountable only to the extent that the formula itself can be inspected and challenged (Busuioc, 2021).

The framework is preserved and the mechanism is replaced. All three components of the leadership architecture are retained, because each answer a question the others cannot. What is discarded is the way each was to be measured: normative statements give way to weighted parameters drawn from binding national regulation. Leadership time management is accordingly reconceptualised from an allocation norm into a Positional Track Record, on the proposition that time invested in becoming a leader is measured not by asserting how hours were divided, but by the verifiable traces those hours left behind.

RESEARCH METHODOLOGY

The study is normative-analytic and design-oriented rather than empirical-explanatory (Salloch et al., 2015). It proceeds in four stages. The first is documentary analysis of the governing instruments: Law 20 of 2023 on the Civil Service; Government Regulations 11 of 2017 and 30 of 2019; Minister of Administrative and Bureaucratic Reform Regulations 38 of 2017 and 3 of 2020 as amended by 20 of 2025; Minister of Finance Regulations 191/PMK.01/2018,

216/PMK.01/2018 and 38 of 2025; and Decree of the Head of the National Civil Service Agency 411 of 2025.

The second stage is formalisation. Each component of the leadership architecture is rewritten as a weighted sum over indicators with explicit conversion scales, and the resulting expressions are ordered into a sequence in which the output of one layer is the input of the next. Rounding, missing-value and range-validation rules are stated so that the instrument is reproducible. The third stage is a traceability audit. Every indicator is mapped to the menu and submenu of Star ASN, the Ministry's internal personnel-record application, and classified as fully available, partially available or unavailable. The proportion of the Final Talent Score resting on unavailable data is then computed analytically for each target level, which converts a qualitative data-quality concern into a quantity that can be reported and tracked.

The fourth stage is simulation. A single officer profile at Echelon III-B is scored against three target-position types, and two sensitivity analyses examine how the Final Talent Score responds to variation in the technical competency test and in certification holdings. This profile is a synthetic construction built to exercise every layer of the instrument under realistic parameter ranges; it does not correspond to, and was not derived from, any actual officer's record. The institutional data used throughout, namely the distribution of 8,353 structural positions across eight echelon tiers and three position types and the corresponding leadership-training completion figures, are drawn from Ministry personnel records at the aggregate level.

Data access for this study was authorised in the author's institutional capacity as Head of Data and Information Systems within the Ministry's Human Resources Development Agency, under whose mandate the Star ASN application and its underlying personnel records fall. All data used in the traceability audit are aggregate counts and distributions, namely positions by echelon tier, position type and training-completion status, rather than individually identifiable personnel records; no name, identification number or individually attributable score appears anywhere in the analysis. The worked simulation described above is, as already stated, an illustrative construction rather than an extract from any individual's file, and this is stated here, at the point where the data source is first declared, rather than only in the limitations that follow.

The author's institutional position also bears on the design itself. The author is an employee of the Ministry whose subject of study is an instrument proposed for that same institution, a relationship common and generally accepted in applied public-administration research, where practitioner access is frequently what makes institution-specific design work possible at all. This is disclosed here as a conflict-of-interest statement: the author has no financial or other interest in the instrument's adoption beyond that ordinary to institutional employment, and the design choices reported below, in particular the weight matrices in Section 4, reflect the author's

professional judgement rather than an outcome the author is positioned to enforce administratively.

The instrument is validated for internal consistency, computability and data feasibility, not for predictive validity. Whether the Final Talent Score predicts subsequent leadership effectiveness is an empirical question that cannot be answered until the instrument has run for several appointment cycles and outcome data have accumulated. Findings are specific to one ministry under one national regime, and the weight matrices in particular reflect deliberate institutional choices rather than empirically derived optima.

RESULT AND DISCUSSION

The Five-Layer Computation Instrument

The instrument is organised as five layers executed in sequence, each consuming the output of its predecessor. The ordering is not arbitrary. Layers one and two evaluate the officer as an individual, independently of any target position, and rest on data that are comparatively complete; an institution beginning implementation can run them immediately. Layers four and five evaluate the officer in relation to a specific target position and are executed only for those who pass the screening gate at layer three. Figure 1 sets out the sequence with its inputs and outputs.

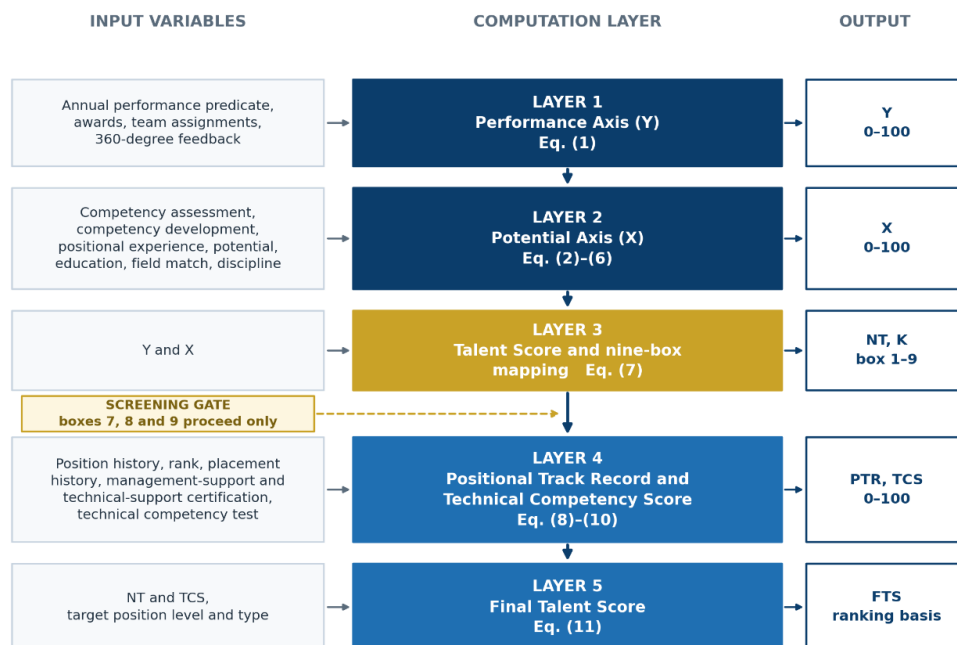


Figure 1 Five-layer structure of the computation instrument, with input variables, governing equations and outputs. PTR denotes the Positional Track Record, TCS the Technical Competency Score and FTS the Final Talent Score

Every layer applies the same algebraic form, which allows the whole instrument to be implemented as a single parameterised routine rather than as a series of special cases. For an indicator set indexed by i with weights w and converted values v :

$$V = \sum_i (w_i \times v_i), \text{ subject to } \sum_i w_i = 1 \text{ and } 0 \leq v_i \leq 100 \quad (1)$$

All indicator values are converted to a common 0–100 scale before weighting, so that every intermediate and final result is itself expressed on that scale and is directly interpretable. Three procedural rules complete the specification. Intermediate results are carried to two decimal places and rounded half-up only at the point of display, so that ranking is never disturbed by rounding order. A missing indicator value is treated as absent rather than as zero: the computation halts and returns a data-incomplete status, because a zero is a substantive claim about the officer whereas an absent record is a claim about the register. Every converted value is validated against the 0–100 range and every weight vector against unit sum before the layer executes.

Layer one, the Performance Axis (Y), is the operational counterpart of leadership work values. The performance parameter comprises Core Performance at 60 per cent, drawn from the annual performance predicate, and Reinforcing Performance at 40 per cent, itself decomposed into three indicators:

$$Y = 0.60 P_{per}^f + 0.15 P_{avv} + 0.15 P_{team} + 0.10 P_{360} \quad (2)$$

where P_{per}^f is the converted annual performance predicate, P_{avv} the award score, P_{team} the team-assignment score and P_{360} the 360-degree feedback score. The weighting of Reinforcing Performance requires a note, because the source regulation is internally inconsistent. Its parameter table specifies awards at 15 per cent, team assignment at 15 per cent and 360-degree feedback at 10 per cent, whereas the worked example in the same annex applies 25 per cent to awards and 15 per cent to team assignment while omitting 360-degree feedback altogether. This instrument follows the parameter table, on two grounds. The parameter table has the standing of a norm while the worked example has the standing of an illustration, so the former governs where they diverge. More substantively, 360-degree feedback is the single indicator in the entire instrument that measures leadership as experienced by those who are led; in a hierarchical organisation with a strong command culture, its removal would leave the measurement of leadership entirely dependent on upward appraisal.

Table 1 Indicators, Weights, Conversion Scales and Data Sources of The Performance Axis

Component	Weight	Indicator	Conversion scale	Data source
Core Performance	60%	Annual performance predicate	Very good 100 · Good 80 · Needs improvement 60 · Poor 40 · Very poor 20	e-Performance service; SKP menu in Star ASN
Reinforcing Performance	15%	Awards received in the last five years	International 100 · National 75 · Inter-institutional 50 · Institutional 25 · None 0	Award history in SIASN and Star ASN
	15%	Team-work assignment in the last two years	Inter-institutional team leader 100 · Internal team leader 75 · Inter-institutional member 50 · Internal member 25 · None 0	Assignment decrees; Other Assignments menu in Star ASN
	10%	360-degree performance feedback	Category score assigned by superiors, peers and subordinates on a 0–100 scale	Not yet available; requires a new module

Source: Compiled from Annex I of Decree of the Head of the National Civil Service Agency 411/2025.

Layer two, the Potential Axis (X), is the operational counterpart of leadership capability. It aggregates four components whose weights sum to unity:

$$X = C + P + Q + I \quad (3)$$

where C is the competency component (40%), P potential (25%), Q qualification (20%), and I integrity and morality (15%). The competency component is the largest and the most composite. It combines a direct assessment of managerial and socio-cultural competency with two proxies for accumulated capability, namely evidence of continuing development and the breadth and depth of positional experience:

$$C = 0.20 A^{comp} + 0.10 D^{comp} + 0.10 E^{pos} \quad (4)$$

$$E^{pos} = (E_{ten} + E_{eiv} + E_{ne}^f) \div 3 \quad (5)$$

Positional experience is the unweighted mean of tenure in the current level (E_{ten}), diversity of positional history (E_{eiv}) and experience in non-definitive assignment (E_{ne}^f). The diversity sub-indicator carries a meaning specific to this Ministry: because it discharges two substantively distinct state functions, an official whose career has crossed from corrections into immigration administration or the

reverse possesses a qualitatively different preparation from one whose career has remained within a single function, and the instrument scores that crossing explicitly.

$$P = 0.25 \times [(\sum_{i=1-8} p_i \div 40) \times 100] \quad (6)$$

The potential assessment scores eight indicators, namely intellectual ability, interpersonal ability, self-awareness, critical and strategic thinking, problem solving, emotional intelligence, learning agility, and motivation and commitment, on a 0–5 scale each, giving a raw total of 0–40 that is rescaled to 0–100.

$$Q = 0.10 Q_e^{du} + 0.10 Q_f^{fd}; I = 0.15 I_i^{ds} \quad (7)$$

Qualification combines formal education level (Q_e^{du}) with the fit of the field of study to the target job family (Q_f^{fd}). Integrity is a single indicator, being verification of the disciplinary record over the preceding five years.

Table 2 Indicators, Weights, Conversion Scales and Data Sources of The Potential Axis

Component	Weight	Indicator	Conversion scale	Data source
Competency	20%	Competency assessment	Managerial and socio-cultural assessment result on a 0–100 scale	Not yet available; requires a new module
	10%	Competency development	Training certificates: $\geq 8 = 100 \cdot 6 - 8 = 75 \cdot 4 - 6 = 50 \cdot 1 - 3 = 25 \cdot \text{none} = 0$. Expertise or tiered certifications: $\geq 3 = 100 \cdot 1 - 2 = 50 \cdot \text{none} = 0$. Averaged.	Training menu in Star ASN
	10%	Positional experience	Tenure: $\geq 5 \text{ years} = 100 \cdot 3 - 4 = 80 \cdot < 2 = 60$. Diversity: cross-institutional = $100 \cdot \text{cross-unit} = 80 \cdot \text{single unit} = 60$. Non-definitive: acting head of region = $100 \cdot \text{acting higher post} = 80 \cdot \text{acting equivalent post} = 60 \cdot \text{temporary higher post} = 40 \cdot \text{temporary equivalent post} = 20 \cdot \text{none} = 0$.	Position and Placement menus in Star ASN
Potential	25%	Potential assessment	Eight indicators scored 0–5, total 0–40, rescaled to 0–100	Not yet available; requires a new module
Qualification	10%	Formal education level	Doctorate 100 · Master 90 · Bachelor/Diploma IV 80 · Diploma III 70 · Upper secondary 60	General Education menu in Star ASN

	10%	Fit of field of study	Consistent with target job family 100 · not consistent 50	General Education menu in Star ASN
Integrity and Morality	15%	Disciplinary record verification	No sanction 100 · light 75 · moderate 50 · severe 25 · currently serving 0	Disciplinary History menu in Star ASN

Source: compiled from Annex I, Tables 1 and 7 to 13, of Decree 411/2025. Job-family fit is read against the immigration, corrections and management-support families of this Ministry

The competency assessment at layer two is confined to managerial and socio-cultural competency. Technical competency is deliberately deferred to layer four, where it is assessed only for officers who have passed the screening gate. The separation matters in an institution of this kind: it permits immigration and corrections expertise to be examined specifically and in depth, using instruments and assessors appropriate to each, without imposing that burden on the initial mapping of the entire workforce.

Layer three, the Talent Score and screening gate. The two axis scores combine in equal proportion into a Talent Score, which simultaneously determines the officer's position on the nine-box map:

$$NT = 0.50 X + 0.50 Y \text{ (8)}$$

Classification applies the categorical bands of the Decree to each axis independently. On the performance axis, a score of 80 to 100 is above expectation, 60 to below 80 meets expectation, and below 60 falls below expectation. On the potential axis the corresponding bands are high, medium and low potential. The three boxes formed by the intersection of the two upper bands on each axis constitute the candidate region, numbered 7, 8 and 9 in the Decree, with box 9 reserved for officers who are simultaneously above expectation and of high potential. Layers four and five execute only for officers within that region, a restriction of substance rather than of computational economy: The Decree places those three boxes in the candidate territory of succession, so technical assessment of officers outside them has no regulatory basis and would generate expectations of advancement that the framework cannot honour.

Layer four, the Positional Track Record and Technical Competency Score, carries the reconceptualised third component. The Positional Track Record measures the investment of leadership time through five verifiable indicators. Positional history registers accumulated leadership experience, since every post held is a period genuinely spent leading. Rank and grade register the productivity of that time, because promotion, and exceptional promotion in particular, is granted

only where elapsed time yielded formally recognised achievement. Placement history registers institutional breadth: The Ministry operates at headquarters, regional-office and technical-implementation-unit level, each with a distinct logic, and an officer who has served at all three has invested time in understanding the organisation from policy formulation through to field delivery. The two certification indicators register deliberate self-development and differ from the other three in temporal direction: where those read time already spent, these read time currently and prospectively invested toward the career being sought.

$$PTR_j = 0.15 R_{his} + 0.10 R_{rnk} + 0.15 R_{pl}^c + w_{ms}(j) M_{ert}^c + w_{ts}(j) T_{ert}^c \quad (9)$$

The subscript j denotes the type of target position. The three fixed indicators total 40 per cent; the two certification weights w_{ms} and w_{ts} vary with j and always sum to the remaining 60 per cent. This is the point at which the structural composition of the Ministry enters the instrument. A prospective work-unit head must command both sides in balance, being answerable for the whole of a unit. A prospective management-support official is required to be stronger in governance, and a prospective technical-support official stronger in substance.

Table 3 Positional Track Record: Conversion Scales and Weights by Type of Target Position

Indicator	Criterion	Value	Work Unit Head	Management Support	Technical Support
Positional history	>10 posts / 6–10 / 1–5	100 / 80 / 60	15%	15%	15%
Rank and grade	IV/b / IV/a / III/d / III/c	100 / 75 / 50 / 25	10%	10%	10%
Placement history	HQ+regional+implementation / HQ+implementation / regional+implementation / HQ+regional / one level only	100 / 80 / 60 / 40 / 20	15%	15%	15%
Management-support expertise	>3 certificates / 1–2 / none	100 / 50 / 0	30%	50%	10%
Technical-support expertise	>3 certificates / 1–2 / none	100 / 50 / 0	30%	10%	50%
Total			100%	100%	100%

Source: Adapted from the internal draft *Indicators and Weights for Competency Assessment of Echelon III Managerial Positions*, Human Resources Development Agency for Immigration and Corrections

The Positional Track Record then enters the Technical Competency Score alongside the technical competency test, in a proportion ρ that the Decree leaves to institutional discretion:

$$TCS = (1 - \rho_l) U + \rho_l PTR_j \quad (10)$$

U is the technical competency test result and ρ_l the proportion allotted to the track record at target level ℓ .

Layer five, the Final Talent Score. The final layer combines the Talent Score with the Technical Competency Score using the level-dependent weight α prescribed by the Decree:

$$FTS(j, \ell) = \alpha_l NT + (1 - \alpha_l) TCS \quad (11)$$

α_l is 0.80 for senior executive, 0.70 for primary executive, 0.60 for administrator and 0.50 for supervisor level, with a permitted reduction of five percentage points where the target position is technically dominant. Because every layer is a weighted sum, the effective contribution of each architecture component to the Final Talent Score can be derived in closed form: leadership capability contributes 0.50α , leadership work values likewise 0.50α , leadership time management $(1 - \alpha)\rho$, and the technical test $(1 - \alpha)(1 - \rho)$. Evaluated across the five leadership levels under the baseline $\rho = 0.20$, this exposes a defect in the design: the effective contribution of leadership time management contracts from 10 per cent at supervisor level to 4 per cent at senior executive level, the opposite of what should occur, since breadth of institutional experience becomes the principal discriminator precisely at the upper levels where competency and performance no longer separate candidates. The full level-by-level breakdown is reported in Appendix A, Table A1.

The remedy does not require departing from the Decree, since ρ is an institutional parameter. Setting ρ to 0.20 for levels one to three, 0.30 for level four and 0.40 for level five stabilises the contribution of leadership time management between 8 and 10 per cent across all levels, leaving the nationally prescribed α untouched. Figure 2 compares the two configurations, and this tiered value of ρ is the one carried forward through the remainder of this article unless otherwise stated.

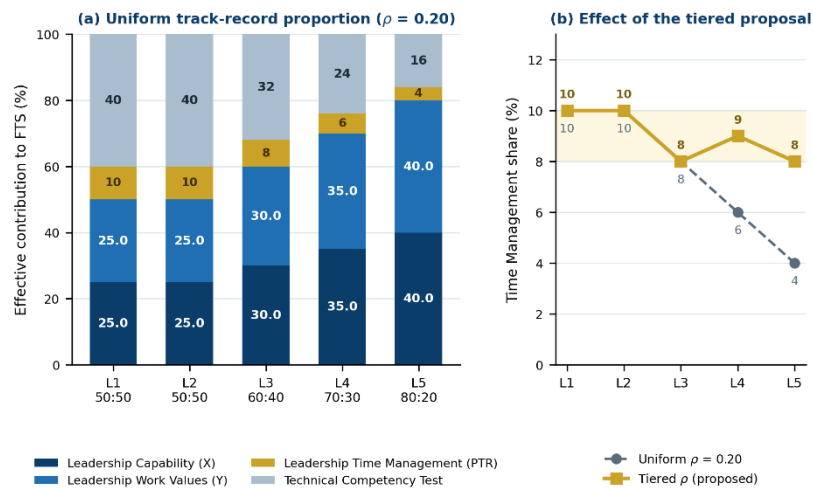


Figure 2 Effective contribution of the four assessment elements to the Final Talent Score. Panel (a) shows the baseline with a uniform track-record proportion; panel (b) shows the stabilising effect of the tiered proportion, with the shaded band marking the 8–10

Traceability Against Star ASN and the Resulting Data Deficit

An instrument that cannot be fed is not an instrument. Star ASN, the Ministry's internal personnel-record application, is organised into six top-level menus: Employee Identity, Employment Track Record, Disciplinary Sanction, Educational Track Record, Publication and Transfer. Every indicator of the instrument was mapped onto that structure and classified as fully available, partially available or unavailable; the complete twenty-row mapping is reported in Appendix B, Table B1. Of the twenty indicator rows traced, six are fully available, ten are partially available and four are unavailable. Partial availability almost always denotes the same defect: the record exists but lacks the classifying attribute the instrument needs. Awards are recorded without a scope code, placements without an institutional-level code, training certificates without a job-family tag. These are schema deficiencies, remediable by adding attributes and back-filling, not by collecting new information. The four unavailable indicators share a more fundamental property: competency assessment, potential assessment, 360-degree feedback and the technical competency test are all direct measurements of the officer, not records of administrative events. Star ASN was designed as a register of what has happened, not as an assessment platform, so their absence cannot be remedied by improving data quality; the assessment instruments themselves must be built.

Because the four unavailable indicators carry known weights, the share of the Final Talent Score resting on data the institution does not have can be computed exactly rather than estimated. Within the potential axis, competency assessment

carries 20 per cent and potential assessment 25 per cent; within the performance axis, 360-degree feedback carries 10 per cent; and each axis enters the Talent Score at 50 per cent. The technical competency test carries the whole of the Technical Competency Score other than the track-record proportion. Hence:

$$\Delta(\ell) = 0.275 \alpha_i + (1 - \alpha_i)(1 - \rho_i) \quad (12)$$

Δ is the share of the Final Talent Score that cannot presently be computed; the coefficient 0.275 is $0.50 \times (0.10 + 0.20 + 0.25)$, the combined weight of the three unavailable assessment indicators within the Talent Score. Equation (12) is evaluated here using the tiered track-record proportion adopted in Section 5.1, namely $\rho = 0.20$ for levels one to three, 0.30 for level four and 0.40 for level five, rather than the uniform baseline of $\rho = 0.20$ used earlier for illustration. At the lower three levels the two parameterisations coincide, so the coverage figures at those levels are unaffected by the choice; at level five they diverge, and this matters for interpretation. Existing data support 46.25 per cent of the Final Talent Score at supervisor level and 66.00 per cent at senior executive level, and the latter figure reflects the tiered value of $\rho = 0.40$: under the unmodified baseline of $\rho = 0.20$, senior-executive coverage would instead be 62.00 per cent. The traceability audit is therefore reported under the corrected parameterisation this article recommends, and this correspondence is stated explicitly here rather than left to be inferred.

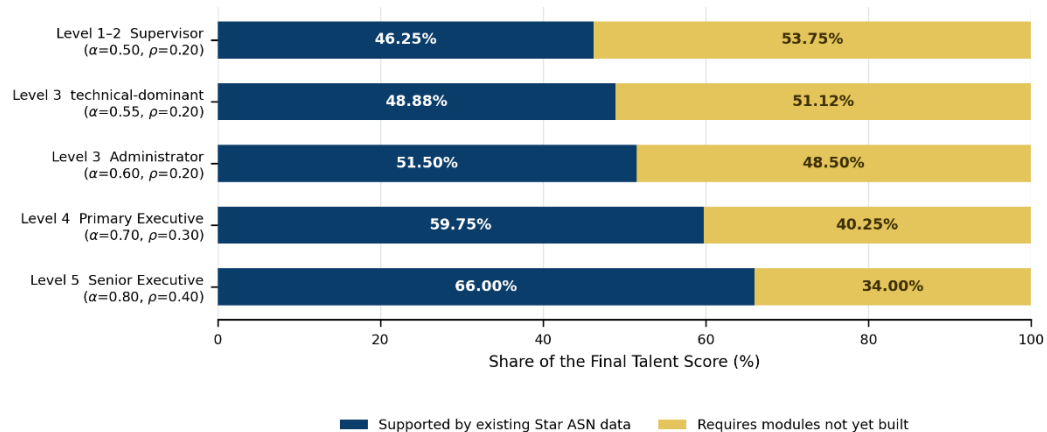


Figure 3 Share of the Final Talent Score supported by data currently held in Star ASN, by target level, computed from Equation (12) under the tiered ρ . The gold segment represents the share dependent on assessment modules that have not yet been built

The pattern is instructive: coverage is weakest precisely where the population is largest. The 7,340 positions at Echelons IV and V, which constitute 87.9 per cent of all structural posts, fall in the band where more than half the score cannot yet be computed. A further caution applies throughout the instrument: certification is invoked twice, once at layer two as a competency-development indicator drawn from the national Decree's own catalogue, and again at layer four as management-

support or technical-support expertise drawn from a Ministry-defined catalogue. A certificate appearing in both registers would be double-counted, inflating the score of officers whose credentials straddle the boundary, so any Ministry catalogue must state explicitly which register each certificate belongs to and the application must enforce that assignment rather than leave it to the data-entry officer.

Two remedial priorities follow, in this order. Attributes must be added to Star ASN, namely a scope code on awards, an institutional-level code on placements, a role-and-scope code on assignments, a job-family tag on training certificates, and a field-of-study code mapped to job families, since these unlock ten partially available indicators at comparatively low cost. Thereafter the four assessment modules must be built, beginning with the technical competency test, which alone carries between 16 and 40 per cent of the Final Talent Score and is the single heaviest missing element at every level below senior executive.

Reference Design for Web-Based Implementation

Pillar 2 of the readiness instrument annexed to Decree 411/2025 requires a web-based talent-management application integrated with the national civil-service information system and with performance services and carries a weight of 20 per cent in the institutional readiness assessment. Pillar 3, weighted at 30 per cent, requires the parameters and weights that the present instrument supplies. The design that follows is therefore not an optional accompaniment to the instrument but the form in which it must be delivered.

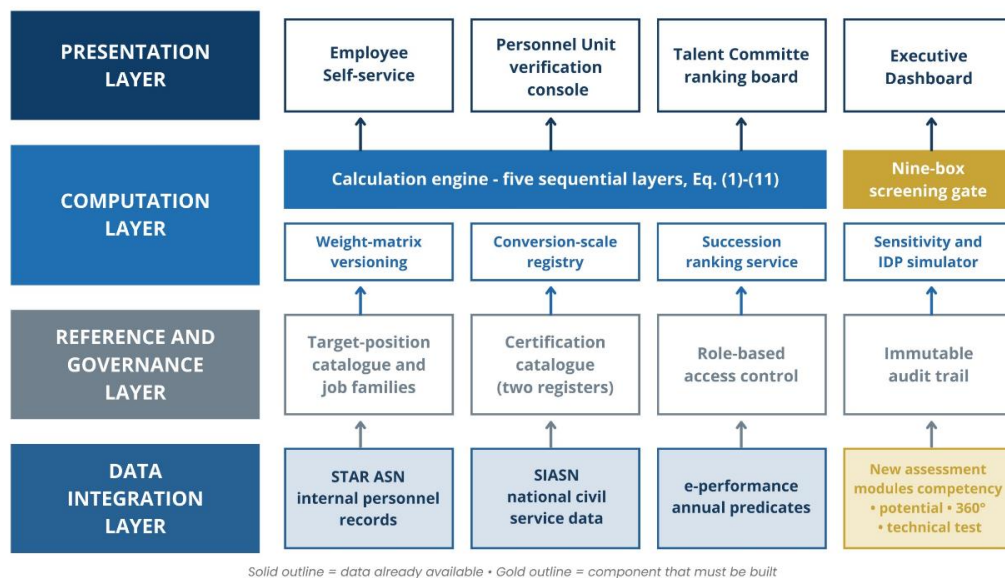


Figure 4 Reference component architecture for the web-based application. Components outlined in gold denote assessment capability that does not yet exist and must be built; all others draw on data already held

The architecture separates four concerns. The data integration layer draws records from Star ASN, the national civil-service information system and

performance services, stamping each field with a completeness marker so that incomplete records are visible before computation rather than after. The reference and governance layer hold the target-position catalogue, the two certification registers, role-based access control and an immutable audit trail. The computation layer executes the five layers in sequence with the nine-box gate interposed. The presentation layer exposes result differently to each role: officers see their own profile and development gaps, personnel units verify source documents, the talent committee sees ranked candidate lists per target position, and executives see aggregate distribution. Ten functional modules implement these concerns in practice, spanning reference-parameter management, data integration, the three new assessment modules (competency, potential and 360-degree feedback), the technical competency test, the calculation engine, succession ranking, dashboards, and simulation for development planning; the full module-by-module specification is reported in Appendix C, Table C1.

Three governance rules deserve emphasis, because omitting them is what causes computed scores to fail under challenge. First, weight matrices must be versioned and every stored result must record the version under which it was produced, so that a later regulatory amendment does not silently invalidate earlier appointment decisions. Second, every figure entering the computation must be traceable to a source document identifier, date and verifying officer; a score that cannot be traced to its evidence will not withstand either an aggrieved candidate or an audit body. Third, the screening gate must be enforced in the engine rather than in the interface, so that technical assessment cannot be initiated for officers outside the candidate region by administrative convenience.

Simulation and Sensitivity Analysis

To establish computability, a single illustrative profile is scored against three target positions at the same level. The officer is an Echelon III-B incumbent with a very good annual performance predicate (100), a national-level award (75), service as leader of an internal team (75) and 360-degree feedback in the good category (80). On the potential axis the officer records a competency assessment of 90, competency development of 50, positional experience of 86.67 (the mean of tenure exceeding five years at 100, cross-unit experience at 80 and acting service in a higher post at 80), a potential assessment total of 38 out of 40, a master's degree (90), a field of study consistent with the target family (100) and an unblemished disciplinary record (100). The track record comprises eight posts held (80), grade IV/a (75), service at all three institutional levels (100), four management-support certificates (100) and one technical-support certificate (50). The technical competency test result is 85.

Table 4 Layer-by-Layer Computation for Three Types Of Target Position

Computation step	Work Unit Head	Management Support	Technical Support	Note
Performance Axis, Y — Eq. (2)	90.50	90.50	90.50	Leadership work values
Potential Axis, X — Eq. (3)	89.42	89.42	89.42	Leadership capability
Talent Score, NT — Eq. (8)	89.96	89.96	89.96	Box 9; passes the screening gate
Positional Track Record — Eq. (9)	79.50	89.50	69.50	Leadership time management; first divergence
Technical competency test, U	85.00	85.00	85.00	A single common test result
Technical Competency Score — Eq. (10)	83.90	85.90	81.90	Baseline proportion $\rho = 0.20$
Weighting $\alpha : 1-\alpha$	60:40	60:40	55:45	Technically dominant variant applied

Source: Author's simulation (2026)

Three observations follow. The first three layers return identical values for all three target positions, because they evaluate the officer as an individual; divergence appears only at layer four, where the target position enters. The same officer consequently obtains three different Final Talent Scores, namely 88.34, 87.54 and 86.33, which is the intended behaviour: succession ranking is specific to the position sought rather than a single universal ordering, and the ordering here reflects a certification profile weighted toward governance. And every figure derives from a record that can be traced to a source document, so the computation is auditable.

The first sensitivity analysis varies the technical competency test while holding all else constant. The partial derivative of the Final Talent Score with respect to the test is $(1 - \alpha)(1 - \rho)$, giving 0.32 on the standard administrator path and 0.36 on the technically dominant path. Each five-point movement in the test therefore shifts the Final Talent Score by 1.60 and 1.80 points respectively.

Table 5 Sensitivity of the Final Talent Score to the Technical Competency Test

Technical competency test	Work Unit Head	Management Support	Technical Support	Deviation from baseline
75	84.34	85.14	82.73	-3.20 / -3.60
80	85.94	86.74	84.53	-1.60 / -1.80
85 (baseline)	87.54	88.34	86.33	—
90	89.14	89.94	88.13	+1.60 / +1.80
95	90.74	91.54	89.93	+3.20 / +3.60

Source: Author's simulation (2026)

Since the spread between adjacent candidates on a ranked list is commonly under two points, a single test administration can determine the order of an entire list. This is the strongest argument for treating construction of the technical competency test as the most urgent item of implementation, and for subjecting it to rigorous quality assurance.

The second analysis varies certification holdings and yields the more consequential result. Where an officer holds more than three certificates in both families, the Positional Track Record reaches 94.50 for all three target types and the weight differentiation disappears entirely. The same collapse occurs when holdings are equal at any level, including zero. The reason is algebraic: the two certification weights always sum to 60 per cent, so when both indicator values are equal, their weighted product is invariant to how that 60 per cent is divided. Differentiation operates only on officers whose certification profile is lopsided.

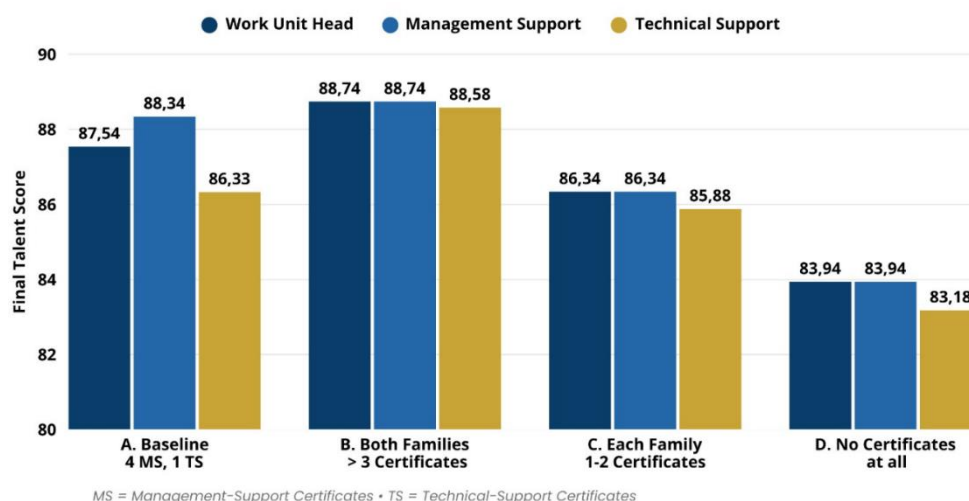


Figure 5 Sensitivity of the Final Talent Score to certification profile across four scenarios and three target position types, with the technical competency test and all other indicators held constant

Differentiated weighting penalises no one. It confers an advantage on officers whose expertise matches the position sought and withdraws that advantage from officers whose expertise is complete on both sides. An officer who builds certification across both families simultaneously attains the highest attainable score on every path without forfeiting any of them. The instrument therefore does not push officers to specialise narrowly; it pushes them to become complete. For an institution short of prepared leaders at every level, that is the incentive worth designing for.

The methodological move at the centre of this article is worth stating in general terms, because the situation that prompted it is not unusual. An institution

adopts a framework from a respected peer; the peer subsequently replaces it; the adopting institution must decide whether to proceed with an instrument its own exemplar has abandoned. The response taken here distinguishes the conceptual contribution of a framework from its measurement mechanism and finds that the two have very different shelf lives. The three questions posed by the leadership architecture, namely what an officer can do, what the officer has demonstrated, and what the officer has invested, remain analytically distinct and individually necessary. What became untenable was the way each was to be scored. Retaining the framework while substituting the mechanism preserves conceptual continuity and gains legal currency at the same time.

This finding extends rather than contradicts the operationalisation-deficit literature reviewed in the Introduction. Vardi and Collings (2023) identify the field's persistent difficulty in moving talent from a binary, either/or construct to a measurable one; the closed-form derivation in Section 5.1, in which every architecture component's contribution to the Final Talent Score can be stated as an explicit function of α and ρ , is one demonstration that the difficulty is tractable once a legally binding measurement regime exists to anchor it. Similarly, Semaihi et al. (2023) find that the link between talent-management practice and individual performance in the public sector depends on mechanisms, such as line managerial support, that generic frameworks leave unspecified; this instrument responds to the same concern by making the mechanism, namely the five-layer computation and its governance rules, the object of specification rather than an implicit assumption.

The reconceptualization of leadership time management is the article's principal substantive contribution and its most contestable one. The original component asked officials to allocate time between managing tasks and managing people. The reformulation measures instead what that time left behind: posts held, grades attained, institutional levels served, certifications earned. The principle is one the measurement literature has long defended, that where a construct cannot be observed directly, proxies carrying verifiable evidence should be preferred over self-reported estimates (Boyatzis, 1991; Spencer & Spencer, 2008). The sensitivity results in Section 5.4 support this choice empirically as well as conceptually: because certification differentiates candidates only when their profiles are lopsided, and because a five-point movement in the technical test can reorder an entire candidate list, the instrument rewards verifiable breadth and depth rather than self-reported allocation, consistent with the succession-planning literature's finding that competency-based rather than narrative succession pipelines are what make leadership transitions defensible (Jackson & Dunn-Jensen, 2021).

The reformulation nonetheless carries an equity risk that should be stated plainly. A track record rewards officers who have been given opportunities. In an organisation where placement across headquarters, regional offices and implementation units is not evenly distributed, and where only 6 per cent of Echelon V incumbents have received the leadership training corresponding to their level,

measuring accumulated experience risks converting historical patterns of access into a formal entitlement. Two features mitigate this. Certification is attainable by individual initiative and carries 60 per cent of the track-record weight, providing a route that does not depend on having been posted. And the component's effective contribution is bounded between 8 and 10 per cent of the Final Talent Score under the tiered proposal, so it modulates rather than determines the outcome. Whether these safeguards suffice is an empirical question that monitoring should answer.

The finding that existing data support only 46 to 66 per cent of the Final Talent Score is presented here as a result rather than a caveat, and the choice is deliberate. Data-quality problems are ordinarily reported qualitatively, in language that resists tracking and accountability. Expressing the deficit as a share of the score, derived in closed form from the instrument's own weights, converts it into a quantity that can be monitored, targeted and reported against. This design choice follows directly from the accountability literature on computed government decisions: an algorithmic or formulaic score is accountable only to the extent that it can be inspected and its shortfalls quantified rather than asserted, and digital government systems of the kind Decree 411/2025 mandates deliver consistency only when institutional discretion operates inside rules that are transparent and auditable (Busuioc, 2021; Criado et al., 2025). The traceability finding also situates this instrument within a wider pattern documented in Indonesian civil-service digitalisation research: a recent evaluation of the SIASN personnel-information system in a regional government agency found that centralised national platforms reliably improve procedural transparency and administrative efficiency while leaving genuine data-interoperability and infrastructure gaps unresolved, a pairing of gains and residual gaps that Star ASN reproduces at the level of individual indicators rather than system architecture (Tambaip et al., 2026). Read against the merit-system literature, the exercise also responds to a broader concern that measurement designs failing to trace to verifiable evidence tend to erode rather than reinforce the government performance they are meant to support Oliveira et al. (2024): quantifying the deficit is what allows this instrument to avoid that failure mode from the outset rather than discovering it after implementation. It also orders the remedial work: because the shortfall concentrates in four assessment-type indicators of known weight, the sequence of remediation follows from the arithmetic rather than from institutional preference. A framework that conceals what it cannot yet compute invites the discretionary substitution it was designed to eliminate.

Four limitations bound the claims made here. The instrument has been validated for internal consistency, computability and data feasibility, not for predictive validity; whether the Final Talent Score forecasts subsequent leadership effectiveness cannot be known until several appointment cycles have produced outcome data. The weight matrices, particularly the three certification vectors, embody deliberate policy choices rather than empirically derived optima, and

calibrating them against realised performance is the most valuable next study. The findings are specific to one ministry operating under one national regime, though the five-layer structure and the traceability method are portable to any institution bound by the same Decree. Finally, the introduction of 360-degree feedback into a hierarchical organisation with a strong command culture requires procedural safeguards, namely guaranteed rater anonymity, defined administration protocols, and protection against retaliation, whose design lies beyond the scope of this article but conditions the validity of a 10 per cent component of the performance axis.

CONCLUSION

This article set out to convert a normative leadership architecture, adopted by the Indonesian Ministry of Immigration and Corrections from a framework its own exemplar institution has since superseded, into a computable instrument capable of supporting defensible appointment decisions. The approach combined documentary analysis of the governing regulations with formal specification: five sequential layers, twelve equations, complete conversion scales, and weight matrices differentiated by target level and position type, together producing a single Final Talent Score for a specified officer-position pair. The three components of the original architecture, namely leadership capability, leadership work values and leadership time management, were preserved conceptually while their measurement mechanism was replaced with parameters drawn from binding national regulation, and the resulting instrument was tested for internal consistency, computability and data feasibility against the institution's own personnel records.

Three findings follow from that test. First, deriving each component's effective contribution to the Final Talent Score in closed form exposed a defect in the baseline design, whereby the contribution of leadership time management contracts as the hierarchy is ascended, and a tiered institutional parameter was proposed that stabilises this contribution between 8 and 10 per cent at every level without disturbing nationally prescribed weights. Second, a traceability audit against the Ministry's personnel-record system found that existing data support between 46.25 and 66.00 per cent of the Final Talent Score, with the shortfall concentrated in four indicators that require new assessment instruments rather than improved record-keeping, and this deficit was expressed as a closed-form share of the score so that remediation can be sequenced and monitored rather than described only in general terms. Third, simulation confirmed that the instrument produces position-specific rather than uniform rankings, and sensitivity analysis established both that a single technical competency test administration can determine the order of an entire candidate list and that differentiated certification weighting rewards officers with complete rather than narrow expertise profiles. A reference architecture for web-based implementation, together with the governance rules a computed score requires to survive challenge, was proposed to carry these findings into practice.

These claims are bounded in four respects, and each points to a direction for further work. The instrument has been validated for computability and data feasibility, not for predictive validity, so establishing whether the Final Talent Score forecasts subsequent leadership effectiveness will require outcome data from several appointment cycles once implementation begins. The weight matrices, particularly the certification vectors governing the Positional Track Record, reflect deliberate institutional policy choices rather than empirically derived optima, and calibrating them against realised performance is the most immediate priority for future research. The findings are specific to one ministry operating under one national talent-management regime, though the five-layer structure and the traceability method are portable to any Indonesian public institution bound by the same Decree and testing that portability in a second institution would strengthen the generalisability of the approach. Finally, introducing 360-degree feedback into a hierarchical organisation with a strong command culture requires procedural safeguards, including guaranteed rater anonymity and protection against retaliation, whose design was outside the scope of this article but whose absence would compromise the validity of a component that carries real weight in the final score.

REFERENCES

- Aeon, B., Faber, A., & Panaccio, A. (2021). Does time management work? A meta-analysis. *PLOS ONE*, 16(1), e0245066. <https://doi.org/10.1371/journal.pone.0245066>
- Asif, A., & Rathore, K. (2021). Behavioral Drivers of Performance in Public-Sector Organizations: A Literature Review. *Sage Open*, 11(1). <https://doi.org/10.1177/2158244021989283>
- Backhaus, L., & Vogel, R. (2022). Leadership in the public sector: A meta-analysis of styles, outcomes, contexts, and methods. *Public Administration Review*, 82(6), 986–1003. <https://doi.org/10.1111/puar.13516>
- Bornay-Barrachina, M., López-Cabrales, Á., & Salas-Vallina, A. (2025). Sensing, seizing, and reconfiguring dynamic capabilities in innovative firms: Why does strategic leadership make a difference? *BRQ Business Research Quarterly*, 28(2), 399–420. <https://doi.org/10.1177/23409444231185790>
- Boyatzis, R. E. (1991). *The Competent Manager: A Model for Effective Performance*. Wiley.
- Budiyanto, A. (2025). THE MEDIATING ROLE OF OCB IN THE RELATIONSHIP BETWEEN COMMITMENT, SATISFACTION, AND PERFORMANCE. *Journal of Multidisciplinary Research*, 4(4), 53–70. <https://doi.org/10.56943/jmr.v4i4.881>
- Busuioc, M. (2021). Accountable Artificial Intelligence: Holding Algorithms to Account. *Public Administration Review*, 81(5), 825–836. <https://doi.org/10.1111/puar.13293>
- Charan, R., Drotter, S., & Noel, J. (2011). *The leadership pipeline : How to build the company*. Jossey-Bass.
- Criado, J. I., Sandoval-Almazán, R., & Gil-Garcia, J. R. (2025). Artificial intelligence and public administration: Understanding actors, governance,

- and policy from micro, meso, and macro perspectives. *Public Policy and Administration*, 40(2), 173–184. <https://doi.org/10.1177/09520767241272921>
- Ezerdi, C., Nurgabdeshev, A., Kozhakhmet, S., Rofcanin, Y., & Demirbag, M. (2022). International HRM in the context of uncertainty and crisis: a systematic review of literature (2000–2018). *The International Journal of Human Resource Management*, 33(12), 2503–2540. <https://doi.org/10.1080/09585192.2020.1863247>
- Jackson, N. C., & Dunn-Jensen, L. M. (2021). Leadership succession planning for today's digital transformation economy: Key factors to build for competency and innovation. *Business Horizons*, 64(2), 273–284. <https://doi.org/10.1016/j.bushor.2020.11.008>
- Koedijk, M., Renden, P. G., Oudejans, R. R. D., Kleygrewe, L., & Hutter, R. I. V. (2021). Observational Behavior Assessment for Psychological Competencies in Police Officers: A Proposed Methodology for Instrument Development. *Frontiers in Psychology*, 12. <https://doi.org/10.3389/fpsyg.2021.589258>
- Le Fouest, S., & Mulleners, K. (2024). Optimal blade pitch control for enhanced vertical-axis wind turbine performance. *Nature Communications*, 15(1), 2770. <https://doi.org/10.1038/s41467-024-46988-0>
- Lušňáková, Z., Dicsérová, S., & Šajbidorová, M. (2021). Efficiency of Managerial Work and Performance of Managers: Time Management Point of View. *Behavioral Sciences*, 11(12), 166. <https://doi.org/10.3390/bs11120166>
- Nzimande, C. S., Ruggunan, S., Maramura, T. C., & Olabiyi, O. J. (2025). Human resource management in the public sector: A study on recruitment, selection and retention practices in South Africa. *SA Journal of Human Resource Management*, 23. <https://doi.org/10.4102/SAJHRM.v23i0.3205>
- Oliveira, E., Abner, G., Lee, S., Suzuki, K., Hur, H., & Perry, J. L. (2024). What does the evidence tell us about merit principles and government performance? *Public Administration*, 102(2), 668–690. <https://doi.org/10.1111/padm.12945>
- Sain, Z. H., Aziz, A. L., Lawal, U. S., Abdullah, N., & Sain, S. H. (2025). EVALUATING THE IMPACT OF DIGITAL TRANSFORMATION ON PUBLIC SERVICE DELIVERY EFFICIENCY AND ACCOUNTABILITY. *Journal of Multidisciplinary Research*, 1–14. <https://doi.org/10.56943/jmr.v4i4.868>
- Salloch, S., Wäscher, S., Vollmann, J., & Schildmann, J. (2015). The normative background of empirical-ethical research: first steps towards a transparent and reasoned approach in the selection of an ethical theory. *BMC Medical Ethics*, 16(1), 20. <https://doi.org/10.1186/s12910-015-0016-x>
- Semaihi, S. O., Ahmad, S. Z., & Khalid, K. (2023). Talent management and performance in the public sector: the mediating role of line managerial support. *Journal of Organizational Effectiveness: People and Performance*, 10(4), 546–564. <https://doi.org/10.1108/JOEPP-09-2022-0274>
- Spencer, L. M., & Spencer, S. M. (2008). *Competence at Work Models for Superior Performance*. Wiley India Pvt. Limited.
- Tambaip, B., Ningrum, R. U., & Muttaqin, M. Z. (2026). Evaluation of new public service delivery through digitalisation of personnel management in local government. *Frontiers in Political Science*, 8. <https://doi.org/10.3389/fpos.2026.1853944>

Vardi, S., & Collings, D. G. (2023). What's in a name? talent: A review and research agenda. *Human Resource Management Journal*, 33(3), 660–682. <https://doi.org/10.1111/1748-8583.12500>